

<https://doi.org/10.26599/FRICT.2026.9441306>

Research Article

U-Net-based hybrid framework for accelerating deterministic rough-surface mixed lubrication simulations

Filimonas Kaliafetis¹✉, Suhaib Ardah¹, James P. Ewen^{1,2}, Daniele Dini¹

¹ Department of Mechanical Engineering, Imperial College London, Exhibition Road, South Kensington, London SW7 2AZ, UK; ² Department of Mechanical Engineering, University of Bath, Claverton Down, Bath BA2 7AY, UK.

✉ Corresponding author. E-mail: f.kaliafetis23@imperial.ac.uk

Received: June 26, 2026; Revised: August 17, 2026; Accepted: August 26, 2026

© The Author(s) 2026.

Abstract:

Despite significant advances in deterministic mixed lubrication modelling, the computational complexity of these simulations remains prohibitive, limiting their practical use in rapid analysis and design optimisation. This study addresses this bottleneck by developing a hybrid machine learning framework that integrates artificial neural networks (ANNs) and convolutional neural networks (CNNs), specifically U-Nets, into a classical finite volume mixed lubrication solver to accelerate rough surface simulations without sacrificing fidelity. The approach uses ANNs to predict smooth pressure,

film thickness, and temperature fields from operating conditions, while U-Nets reconstruct rough fields from smooth solutions and rough surface topographies. The ANNs achieve high accuracy for smooth fields, with mean R_{rs}^2 values of 0.998 for pressure, 0.999 for film thickness, and 0.925 for temperature. The U-Nets generalise well to rough surfaces, achieving R_{rs}^2 values of 0.990, 0.996, and 0.938 for pressure, film thickness, and temperature respectively. Six hybrid configurations are assessed, all of which reduce computational cost relative to the classical solver (average simulation time: 10.36 minutes). The most efficient framework delivers a 74.39% reduction in computation time across all testing samples. Overall, the results demonstrate that hybrid machine learning-physics frameworks can substantially accelerate deterministic mixed lubrication simulations while preserving high predictive accuracy.

Keywords:

Thermo-elastohydrodynamic lubrication, Mixed lubrication, Surface roughness, Convolutional neural networks, U-Nets

1. Introduction

Thermo-elastohydrodynamic lubrication (TEHL) simulations have long been used to accurately model tribological contacts in bearings, gears and other components and mechanisms. TEHL theory extends classical lubricated contact modelling by coupling hydrodynamic fluid film behaviour and elastic deformation of solid bodies with thermal effects arising from viscous dissipation and material conductivity. Traditional TEHL simulations solve a set of nonlinear partial differential equations, including the Reynolds equation, elasticity equations and energy equations, to predict pressure, film thickness and temperature profiles, among others, with high fidelity [1].

A central challenge in TEHL modelling is the accurate representation of real surface topography. For idealised smooth surfaces, the Reynolds equation predicts a continuous pressurised fluid film that entirely separates the contacting solids. However, in practice, surface roughness often has characteristic scales that are comparable, or even exceed, the film thickness predicted for smooth surfaces, particularly under severe operating conditions. Moreover, in recent years there has been a significant push towards lower viscosity lubricants, particularly in the automotive sector, to minimise power losses and improve energy efficiency [2]. Although lower viscosity lubricants reduce viscous friction losses, they also tend to produce thinner lubricant films, increasing the likelihood of mixed lubrication conditions in which asperity interactions occur. Under these conditions, the accurate modelling of roughness effects in tribological contacts becomes more important than ever.

Substantial research effort has been devoted to the development of deterministic mixed lubrication models capable of explicitly resolving surface roughness effects representative of real engineering surfaces. Unlike stochastic roughness approaches, deterministic models directly incorporate measured or generated surface topography into the numerical solution of the lubrication problem [3]. Ai and Cheng [4] used a deterministic approach to model moving surface defects in heavily loaded point contacts with good agreement against experimental results. Chang [5] used a deterministic model to simulate the contact of two colliding asperities as well as the contact between two rough surfaces with sinusoidal roughness profiles in the mixed lubrication regime. Hu and Zhu [6] employed a unified deterministic model to simulate various lubrication conditions, from full-film elastohydrodynamic lubrication down to boundary lubrication in point contacts. Their approach was based on the idea that the same Reynolds equation may be used for both the hydrodynamic and contact regions. Their model proved to be able to provide solutions even at very low film thickness ratios (λ). Subsequent studies further extended deterministic approaches to include thermal effects, transient operating conditions and more complex surface geometries [3], [7].

Substantial advances have also been made in the numerical methods used in lubrication models. Efficient discretisation strategies, multigrid methods, fast Fourier transform techniques and advanced iterative solvers have significantly improved solution robustness and convergence behaviour [8], [9], [10], [11]. Nevertheless, the complexity of such simulations means that the computational burden remains significant [12]. The multiscale nature of rough surface lubrication problems requires fine spatial discretisation in order to accurately resolve asperity-scale interactions while simultaneously capturing the overall contact behaviour. In addition, the inclusion of thermal effects introduces further nonlinear coupling between temperature-dependent lubricant properties, heat generation and hydrodynamic effects. Furthermore, the inclusion of transient effects compounds these challenges, as the above calculations have to be repeated for individual time iterations. As a result, high-fidelity deterministic lubrication simulations usually require substantial computational resources and long execution times, limiting their applicability real-time engineering applications such as digital twins and online condition monitoring systems.

Machine learning models have been rapidly emerging as a way to accelerate such simulations, due to their near-instantaneous prediction capability once trained. More specifically, machine learning has been used to predict friction coefficients, wear rates, film thickness and pressure distributions, as well as to develop surrogate models for elastohydrodynamic and thermo-elastohydrodynamic lubrication problems [13], [14], [15], [16], [17]. Physics-informed neural networks (PINNs) have also been explored in this context. These embed governing physical equations directly into the learning process, enabling the solution of lubrication-related partial differential equations without fully explicit numerical discretisation [18], [19], [20].

However, the majority of existing machine learning approaches in lubrication modelling have focused on idealised or nominally smooth contact conditions, where surface topography is either neglected or represented in a simplified manner. Comparatively limited attention has been given to the application of machine learning techniques to rough surface simulations, particularly in the context of deterministic TEHL modelling where multiscale surface features significantly influence the tribological response of the system. In addition, while PINNs offer a physics-driven alternative to purely data-driven models, their application to elastohydrodynamic lubrication remains challenging due to the strong coupling between pressure and elastic deformation, as well as difficulties in resolving steep gradients, enforcing boundary conditions, and capturing complex multiscale behaviours [21]. Capturing rough surface interactions in general remains particularly challenging due to the highly nonlinear coupling between roughness-induced pressure perturbations, local film thickness variations and thermal effects. Nevertheless, recent developments indicate that there is rapid progress in this area. In a recent study by Zhao et al. [22], the authors developed an Adaptively Weighted Multi-Scale PINN framework that was able to accurately solve EHL problems for smooth, periodically rough as well as fractal rough surfaces. However, the model was trained and evaluated for specific operating conditions, limiting its general applicability across varying loads and speeds.

Another type of neural networks, convolutional neural networks (CNNs), have also been used in applications where spatially varying fields have to be resolved and interpreted. CNNs contain convolution kernels that are used to extract spatial features and detect patterns in the inputs [23]. The pooling layers reduce the resolution of the feature maps (outputs of the convolutional layers) to achieve shift-invariance, which allows the network to recognise features regardless of their actual location in the spatial domain. However, a disadvantage of classic CNNs is that the downsampling process in the pooling layers can lead to significant detail loss in the output [24]. To solve this problem U-Nets can be used.

U-Nets are a type of CNN typically extensively used in medical applications, primarily for image segmentation [25], [26], [27]. These have an encoder-decoder type architecture. The encoder block is similar to a classic CNN with convolution and pooling layers that downsample the feature maps to extract high-level features. The decoder block is used to increase the resolution of the feature maps, restoring any detail lost. Skip connections link corresponding encoder and decoder layers, allowing fine-grained details lost during downsampling to be reintroduced into the decoder, improving the reconstruction of high-resolution outputs [28]. This architecture enables U-Nets to capture both global and local features, which is particularly important in applications such as tribology, where microscopic surface roughness interacts with macroscopic pressure, film thickness and temperature fields.

Overall, existing machine learning approaches to lubrication modelling remain limited in several key respects. For example, the majority have been developed for smooth or idealised contact conditions rather than deterministic rough surfaces, and in cases where rough surfaces have been considered, models have typically been restricted to specific operating conditions with limited generalisability. Furthermore, the majority of existing machine learning surrogate models in lubrication research aim to fully replace the numerical solution process, rather than being integrated within it. Hybrid approaches that couple machine learning predictions with a classical iterative solver remain comparatively unexplored, having previously only been investigated by the authors for smooth TEHL problems [12]. In addition, several existing machine learning studies in lubrication modelling focus on the prediction of pressure and/or film thickness alone, with comparatively less attention given to the coupled thermal field, despite its critical role in lubrication performance.

The aim of this work is, therefore, to develop and evaluate a deep learning framework capable of predicting rough-surface mixed lubrication field distributions from corresponding smooth-field

solutions and surface topography data. Specifically, a U-Net architecture is employed to learn the nonlinear spatial mapping between smooth pressure, film thickness, and temperature profiles, combined with surface roughness information, and the resulting rough-field responses. By doing so, this study seeks to provide a computationally efficient surrogate model that retains spatial resolution while significantly reducing the cost associated with deterministic rough surface mixed lubrication simulations.

2. Methodology

2.1. Unified lubrication framework

The details of the classical mixed lubrication framework used in this study are presented in [12] and [29]. In summary, the model is based on a unified finite volume formulation capable of resolving TEHL and mixed lubrication behaviour in both smooth and deterministic rough surface contacts. The framework solves the coupled generalised Reynolds and energy equations, while also accounting for mass-conserving cavitation, elastic deformation of the contacting bodies, and transient thermal effects in the lubricant and solids. A deterministic surface roughness description is incorporated directly into the film thickness, allowing asperity-level interactions and transition between full-film and asperity contact regimes to be captured within the same numerical formulation. The lubricant rheology is described using a pressure- and temperature-dependent viscosity model based on the Roelands equation, while density variations are captured using the Dowson–Higginson relationship. In addition, shear-thinning behaviour is accounted for using an Eyring-type non-Newtonian model. These constitutive relations enable accurate representation of lubricant behaviour under the high-pressure, high-shear, and thermally coupled conditions characteristic of mixed lubrication. The solution is obtained using a fully coupled iterative scheme with appropriate convergence criteria for pressure, temperature, and load balance, as described in detail in the above references.

A flowchart of the classical framework is shown in Figure 1. The total pressure (hydrodynamic and asperity), cavitation, film thickness and deformation profiles are calculated within the unified hydrodynamic solver (red box). The temperatures profiles of the fluid and solid bodies are calculated within the thermal solver (blue box). Within each global iteration (green box), the hydrodynamic and thermal solvers are run sequentially. This process is repeated until the residual changes in the pressure and temperature profiles are small enough to satisfy the global convergence criteria, at which point the solver is considered converged for that time iteration. The term *global iteration* used throughout the rest of this paper refers to a single pass through the coupled hydrodynamic-thermal loop within a given time iteration.

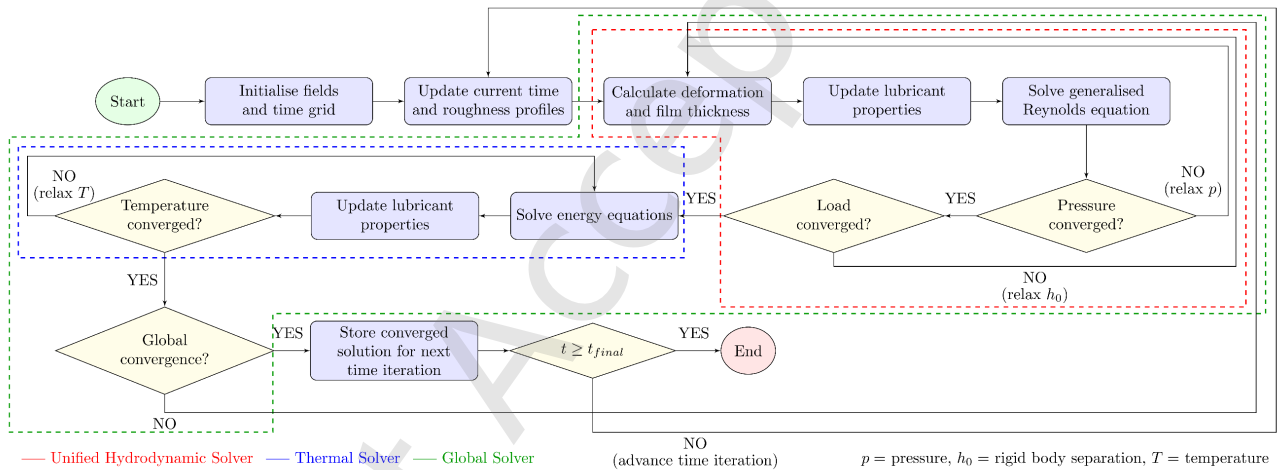


Figure 1 – Unified mixed lubrication framework solution procedure.

2.2. Acceleration of mixed lubrication model

In rough surface mixed lubrication simulations, the smooth solution is normally computed first, with the roughness profiles gradually introduced into the domain over a few time steps (e.g. 30-50 time steps according to Wang and Zhu [1]) to ensure stability. The fact that the solutions for many time steps have to be computed sequentially makes rough surface simulations time consuming. In this study, it is investigated whether the subsequent time steps can be skipped altogether, with a final solution obtained using a combination of the smooth profiles only and machine learning models.

In a previous study [12], artificial neural networks (ANNs) were used to predict key parameters in tribological systems that were in turn used as initial predictions for the mixed lubrication models. In that study only smooth surfaces were considered and the inputs to the ANNs were purely scalar values, like load or entrainment speed. In this paper, that approach is extended to rough surfaces.

Predicting the results of rough surface mixed lubrication simulations with ANNs that can only accept scalar values is impossible, since two rough surfaces can be very different even if they have the same root mean square (RMS) roughness, skewness and kurtosis, all of which are parameters used to characterise such rough surfaces. Therefore, conventional fully connected ANNs are unsuitable in this scenario. Convolutional neural networks (CNNs), on the other hand, are more suitable since they are able to accept images, or matrices, as inputs and are particularly useful for recognising spatial patterns [23].

In the proposed framework, the CNNs accept the pressure, film thickness, velocity and temperature profiles from the smooth solution, as well as the final rough surface profiles, as inputs, and in turn output the rough profiles of pressure, film thickness and temperature, either skipping the rough iterations altogether or using such predictions to accelerate the rough iterations. ANNs are also used to accelerate the smooth iteration of the solver in the manner described in [12]. A flowchart of the proposed hybrid framework is presented in Figure 2. Both the smooth and rough iterations are executed within the classical framework described previously.

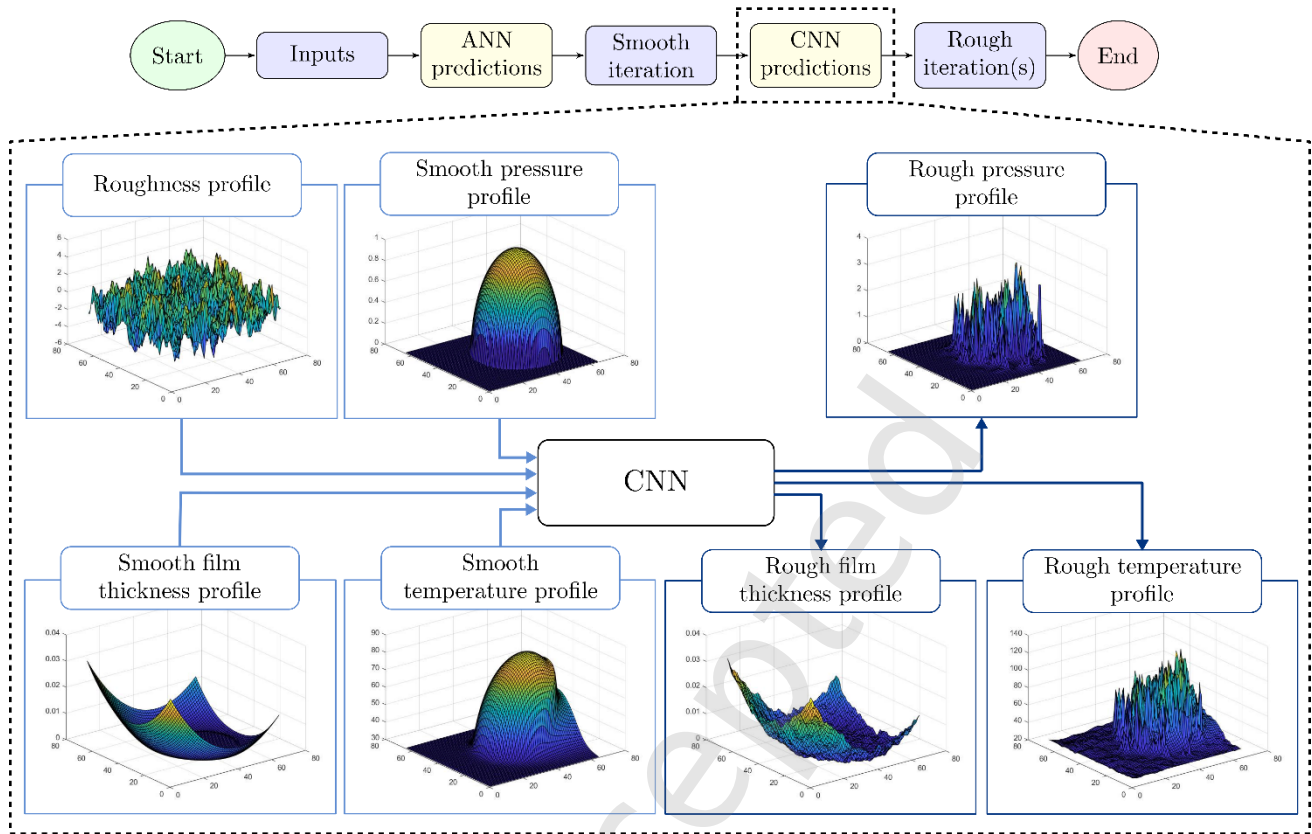


Figure 2 – Hybrid framework flowchart.

More specifically, six different hybrid frameworks are compared. In all of them, as well as the classical finite volume framework (Classical FW), only two time iterations are simulated. In the first time iteration the surfaces are smooth and in the second time iteration the fully developed roughness profiles are introduced into the simulation. It should be noted that this two-time-iteration structure is not intended to represent a practical transient simulation, in which roughness would instead be introduced gradually over many sequential time steps. Rather, it is used here to provide a direct and consistent basis for comparing the accuracy and computational cost of the Classical FW against the proposed hybrid frameworks.

In Hybrid FW 1, 2 and 3 the first time iteration is run in the same way as the Classical FW, at the end of which the CNN predictions take place. The differences between the frameworks are evident in the

second time iteration. In Hybrid FW 1, after the CNN predictions take place, the second time iteration is allowed to run to completion. In Hybrid FW 2, the second time iteration is limited to just one global iteration. In Hybrid FW 3, the second time iteration is skipped altogether and the CNN predictions make up the final solution. In this way, varying amounts of acceleration and solution accuracy can be achieved. The final three hybrid frameworks, Hybrid FW 1 (+ANN), Hybrid FW 2 (+ANN), Hybrid FW 3 (+ANN), are identical to Hybrid FW 1, 2 and 3, respectively, but also include ANN predictions for smooth pressure, film thickness and temperature profiles at the beginning of the first iteration, with the first iteration being allowed to run to completion in the same manner as [12]. Table 1 shows a summary of the six different hybrid frameworks.

Table 1: Hybrid simulation frameworks.

Framework	ANN prediction	Smooth iteration	CNN prediction	Rough iteration
Hybrid FW 1	No	Same as Classical FW	Yes	Same as Classical FW
Hybrid FW 2	No	Same as Classical FW	Yes	Limited to one global iteration
Hybrid FW 3	No	Same as Classical FW	Yes	Skipped
Hybrid FW 1 (+ANN)	Yes	Same as Classical FW	Yes	Same as Classical FW
Hybrid FW 2 (+ANN)	Yes	Same as Classical FW	Yes	Limited to one global iteration
Hybrid FW 3 (+ANN)	Yes	Same as Classical FW	Yes	Skipped

To quantify the amount of acceleration of each Hybrid FW simulation, the percentage difference is calculated between the Hybrid FW simulation time and the Classical FW simulation time. The simulation time is defined as the wall-clock time elapsed between the inputs being defined and global convergence being reached and does not include any post-processing like file-saving. All simulations

were performed using MathWorks MATLAB R2024a on a computer equipped with a 13th Gen Intel® Core™ i7-13850HX (base speed 2.10 GHz), 64 GB of RAM and running the Windows 11 Enterprise operating system.

2.3. Training and testing data generation

In this study, both the operating conditions and individual roughness profiles are varied across the training and testing samples. These operating conditions varied include the applied load, entrainment speed and slide-to-roll ratio. The range of those values are chosen to reflect the range of operating conditions of a mini traction machine (MTM) [30]. The solid body and lubricant properties remain constant across all samples. Similarly to [12], the Latin Hypercube Sampling (LHS) method was used to sample data points with the desired range of values using the MATLAB function *lhsdesign*. The generated samples were categorised using the dimensionless load (M) and material (L) Moes parameters. The expressions for M and L for point contacts are given below [1]:

$$M = W^*(2U^*)^{-3/4} \quad (1)$$

$$L = G^*(2U^*)^{1/4} \quad (2)$$

where $W^* = \frac{W}{E'R^2}$, $U^* = \frac{\eta_0 u_{ent}}{E'R}$ and $G^* = \alpha_P E'$.

Table 2 shows the operating conditions varied across the samples and Table 3 shows the parameters remaining constant across all simulations.

Table 2: Training and testing sample parameters.

Parameter	Min.	Max.
Applied load, W [N]	10	100
Entrainment speed, u_{ent} [m/s]	0.01	3

Parameter	Min.	Max.
Slide-to-roll ratio, SRR [-]	0.1	2

This study is divided into three phases. In Phase 1, Gaussian isotropic rough surfaces with identical statistical roughness parameters were used for training and testing. Although the statistical properties were fixed, the specific roughness profiles differed between samples due to their random generation. In Phase 2, the networks trained in Phase 1 were evaluated on Gaussian rough surfaces with varying Peklenik number (γ) values in order to assess their ability to generalise across different surface anisotropies. Finally, in Phase 3, the trained networks were tested on Gaussian rough surfaces with varying root mean square (RMS) roughness (R_q) values to evaluate their sensitivity to changes in roughness amplitude. In this way, the predictive capability of the neural networks can be assessed both for interpolation within the trained roughness distribution and for extrapolation to surfaces with modified statistical characteristics. Training was deliberately restricted to a single, fixed roughness distribution in Phase 1 so that Phases 2 and 3 could evaluate whether the networks generalise to previously unseen anisotropies and roughness amplitudes without the need of retraining.

Table 3: Constant parameters.

Parameter type	Parameter	Value
Operating conditions	Ambient temperature, T_{amb} [$^{\circ}\text{C}$]	30
	Lubricant reference viscosity, η_0 [$\text{Pa}\cdot\text{s}$]	0.05
Lubricant properties	Lubricant pressure-viscosity coefficient, α_p [$1/\text{Pa}$]	2×10^{-8}
	Fluid specific heat capacity, c_p [$\text{J}/(\text{kg}\cdot\text{K})$]	2000
	Fluid thermal conductivity, k [$\text{W}/(\text{m}\cdot\text{K})$]	0.14
	Fluid coefficient of thermal expansion, β [$1/\text{K}$]	6.5×10^{-4}
	Eyring shear stress, τ_E [Pa]	1×10^7
	Reference density, ρ_0 [kg/m^3]	875
Solid properties	Effective radius of curvature, R [m]	0.01

Parameter type	Parameter	Value
Simulation parameters	Young's modulus of solid bodies, $E_{1,2}$ [Pa]	2×10^{11}
	Poisson's ratio of solid bodies, $\nu_{1,2}$ [-]	0.3
	Thermal conductivity of solid bodies, $k_{1,2}$ [W/(m·K)]	21
	Specific heat capacity of solid bodies, $c_{p1,2}$ [J/(kg·K)]	470
	Density of solid bodies, $\rho_{1,2}$ [kg/m ³]	7850
	Boundary friction coefficient, f_b [-]	0.1
	Mesh size of fluid and solid domains, $N_x \times N_y \times N_z$	$64 \times 64 \times 11$
	Computational domain	$-1.9 \leq x/a \leq 1.1,$ $-1.5 \leq y/a \leq 1.5$
	Pressure, liquid film fraction convergence criterion, $e^{p,\theta}$	1×10^{-5}
	Load convergence criterion, e^W	1×10^{-4}
	Temperature convergence criterion, e^T	1×10^{-4}
	Global pressure convergence criterion, e_{global}^p	1×10^{-3}
	Global temperature convergence criterion, e_{global}^T	1×10^{-3}
	Number of time iterations, N_t	2
	Time step magnitude, Δt [s]	$\Delta t = \Delta x / u_{ent}$

In Phase 1, 2000 samples were generated, with 85% of them (1700 samples) used for training and 15% (300 samples) used for testing the accuracy of the predictions. Phases 2 and 3 consisted of 75 testing samples each. A summary of the roughness parameters of the surfaces used in the three phases is presented in Table 4.

Table 4: Roughness parameters.

Parameter	Phase 1	Phase 2	Phase 3
Root mean square roughness, $R_{q,1,2}$ [m]	1×10^{-7}	1×10^{-7}	$1 \times 10^{-8} - 1 \times 10^{-7}$
Skewness, $R_{sk,1,2}$ [-]	0	0	0
Kurtosis, $R_{ku,1,2}$ [-]	3	3	3
Correlation length in x-direction, $l_{x,1,2}^*$ [m]	5×10^{-5}	5×10^{-5}	5×10^{-5}

Parameter	Phase 1	Phase 2	Phase 3
Peklenik number, $\gamma_{1,2} = l_{x,1,2}^*/l_{y,1,2}^* [-]$	1	0.1 – 10	1

While parameters such as RMS roughness and the Peklenik number provide convenient scalar descriptors of surface amplitude and anisotropy, they do not fully characterise the spatial frequency content of the surface. For this reason, the surface topographies are further described in terms of their power spectral density (PSD), which provides a frequency-domain representation of roughness and distinguishes contributions from different lengths scales.

The importance of PSD-based characterisation has been highlighted in recent surface metrology studies, which emphasise that conventional roughness parameters alone are insufficient to uniquely define surface topography. In particular, the Surface-Topography Challenge study [31] demonstrated the necessity of spectral analysis for meaningful comparison and reproducibility of rough surfaces. Accordingly, PSD spectra of the generated surfaces in the training and testing samples are presented in Section 3.1 to provide a more complete description of their multiscale characteristics. These were calculated using the method described in [31].

2.4. Machine learning models

2.4.1 Artificial neural networks

In this study, ANNs were used for the prediction of the smooth profiles of pressure, film thickness and temperature. Their use is motivated by previous work by the authors [12], in which they demonstrated strong predictive capability for smooth TEHL problems using scalar operating parameters as inputs.

Three separate ANNs were constructed in MATLAB for smooth pressure, film thickness and temperature field predictions. The ANNs consisted of multiple hidden layers, each made up of a

fullyConnectedLayer, a *swishLayer* and a *dropoutLayer* to avoid overfitting. The scalar inputs were passed to the ANNs using a *featureInputLayer* with a feature size equal to the number of input parameters (three). A final *fullyConnectedLayer* with size equal to the number of output values of each ANN is used in addition to a *regressionLayer* to handle the outputs. The adaptive moment estimation (ADAM) algorithm is used for training the ANNs as was done in [12].

Due to the small number of scalar inputs and their comparable magnitudes across the sampled operating range, additional input normalisation was not required in the present study. This differs from the approach adopted in [12], where normalisation was necessary due to the inclusion of additional parameters spanning several orders of magnitude. For example, Young's modulus varied between 70×10^9 and 400×10^9 Pa while the lubricant pressure–viscosity coefficient ranged between 1×10^{-8} and 2.5×10^{-8} Pa⁻¹.

2.4.2 Convolutional neural networks

In this study, the U-Nets were constructed in MATLAB using the built-in *unet* function. Three separate U-Nets were constructed for rough pressure, film thickness and temperature field predictions. The inputs were passed to the U-Nets using an *imageInputLayer* allowing the spatial fields to be treated as two-dimensional images. The profiles from the smooth solution used as inputs to the U-Nets implicitly capture the variations within the training and testing space, as changes in operating conditions are directly reflected in their spatial distributions. These inputs are summarised in Table 5.

Table 5: U-Net inputs.

U-Net	Input	Normalisation	Type
All	Target roughness profile	$(s_1 + s_2)R/a^2$	One 2D matrix ($N_x \times N_y$)
	Smooth pressure profile	p/p_{Hertz}	One 2D matrix ($N_x \times N_y$)

U-Net	Input	Normalisation	Type
	Smooth temperature profile	$\ln(T/T_{amb})$	Multiple 2D matrices $N_z \times (N_x \times N_y)$
	Smooth fluid velocity profile in x-direction	u_f/u_{ent}	Multiple 2D matrices $N_z \times (N_x \times N_y)$
Pressure / temperature	Smooth film thickness profile	hR/a^2	One 2D matrix ($N_x \times N_y$)
Film thickness	Smooth film thickness profile	$(h - h_{macro})R/a^2$, where $h_{macro} = \frac{x^2}{2R} + \frac{y^2}{2R}$	One 2D matrix ($N_x \times N_y$)

The inputs to the pressure and temperature U-Nets were identical. However, the smooth film thickness profile input to the film thickness U-Net was modified slightly by removing its macroscale geometry component. This was done because the macroscopic film thickness variation has a significantly larger magnitude than the roughness-induced perturbations within the Hertzian contact region. As a result, that macroscale component can dominate the loss function during training and reduce the network's sensitivity to the smaller-scale variations that are of primary interest. By subtracting the macroscale geometry, the network was encouraged to focus on learning the roughness-induced deviations in film thickness within the contact region, improving predictive accuracy where it is most critical.

This modification was not applied to the pressure and temperature networks, as the rough pressure and temperature fields are predominantly governed by their corresponding smooth distributions combined with the surface roughness effects. In those cases, the smooth field provides the appropriate large-scale structure for the prediction, and its magnitude does not mask the local roughness-induced variations in the same manner observed for film thickness.

2.4.3 Hyperparameter optimisation

To ensure optimal network performance, the hyperparameters of the ANNs were systematically tuned using Bayesian optimisation. Bayesian optimisation iteratively evaluates different hyperparameter combinations and uses the results of previous iterations to guide the selection of new candidates, focusing the search on regions of the hyperparameter space that are likely to improve performance [32]. This approach is more efficient than random or grid search methods and is particularly useful when training neural networks is computationally expensive.

The objective of the optimisation was to maximise the predictive accuracy of each network, represented by the mean rescaled coefficient of determination (R_{rs}^2) of the testing dataset across all samples. The parameter R_{rs}^2 is defined as

$$R_{rs}^2 = \frac{1}{2 - R^2}, \quad (3)$$

where R^2 is the conventional coefficient of determination

$$R^2 = 1 - \frac{\sum_i (y_i - \hat{y}_i)^2}{\sum_i (y_i - \bar{y})^2}, \bar{y} = \frac{1}{N} \sum_{i=1}^N y_i. \quad (4)$$

Here, y_i is the actual value of a variable, $\hat{y}_i$ is the predicted NN value and N is the number of output values of each NN for each sample. For example, for 2D profiles (pressure and film thickness) $N = N_x \times N_y = 64 \times 64 = 4,096$ and for 3D profiles $N = N_x \times N_y \times N_z = 64 \times 64 \times 11 = 45,056$. The use of a rescaled coefficient of determination ensures that the performance metric remains bounded between 0 and 1, avoiding the misleading interpretation of negative R^2 values.

For the ANNs, the optimisation explored the number of hidden layers, the number of neurons per hidden layer, and the dropout rate. For the U-Nets, the encoder depth, which determines the number of times the input image is downsampled and upsampled, was varied from 3 to 5, while the number of filters in the first encoder layer were varied from 16 to 64. Unlike the ANNs, these hyperparameters

were explored manually by repeating training runs with different configurations and evaluating performance. The full ranges of hyperparameters considered for both ANNs and U-Nets are provided in Table 6.

Table 6: Hyperparameter ranges.

Network	Parameter	Range
ANN	Number of hidden layers	1-5
	Number of hidden layer neurons	16-512
	Dropout layer probability	0.01-0.5
U-Net	Encoder depth	3-5
	Number of first encoder filters	16-64

3. Results and discussion

3.1. Training and testing data

Figure 3 shows the parameter space covered by the training and testing samples in this study. The Moes dimensionless parameters for the dataset range from about 5 to 15 for M and from about 10 to 1100 for L . Since all values of L are significantly greater than unity, the simulations lie entirely within the piezoviscous regime. The selected range of M corresponds to elastic behaviour rather than the rigid asymptotic limit. Consequently, all samples fall within the elastic–piezoviscous regime as classified by Moes [33] and later by Nijenbanning et al. [34]. This regime is representative of practical EHL contacts encountered in engineering applications.

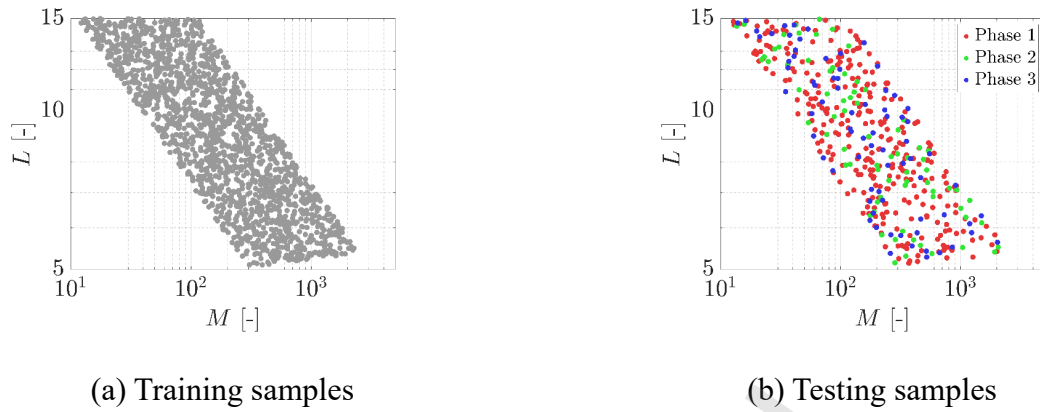
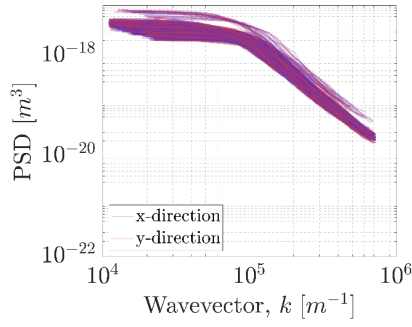


Figure 3 – Distribution of all training and testing samples in the Moes parameter space, defined by the dimensionless load parameter (M) and speed parameter (L).

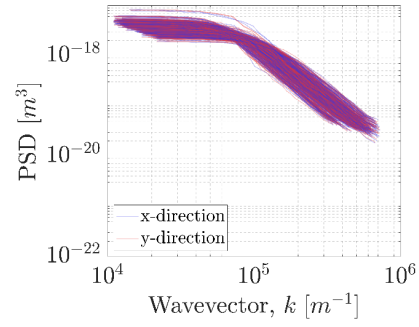
Figure 4 highlights the PSD distributions of the rough surfaces used in the training and in the three testing phases. The PSDs are separated in the x- and y-directions to highlight anisotropic effects. For the training and Phase 1 testing samples, the PSD curves largely overlap. This is expected since both datasets consist of Gaussian isotropic surfaces with identical statistical roughness parameters. Although the individual profiles differ due to random generation, their spectral content remains statistically consistent, resulting in similar distributions in both spatial directions.

In Phase 2, where the Peklenik number (γ) is varied, a clear change in anisotropy is observed. Since the correlation length in the x-direction is kept constant, the PSD curves in the x-direction remain largely unchanged compared to Phase 1. However, the correlation length in the y-direction does vary, leading to a noticeable spreading of the PSD curves along that direction. This behaviour confirms that in Phase 2, anisotropy is introduced while preserving the underlying roughness amplitude.

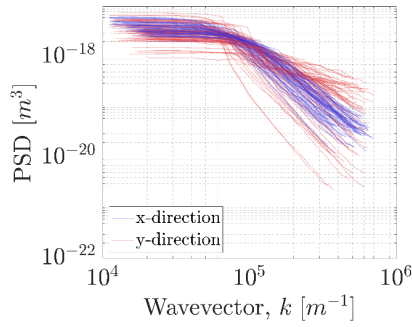
Lastly, in Phase 3, where the RMS roughness is varied, the overall shape of the PSD curves remains similar to Phase 1, indicating that the spatial frequency content is preserved. However, the vertical spread of the PSD curves is increased due to the variation in RMS roughness. This is expected, since the PSD amplitude scales with the square of surface height [31].



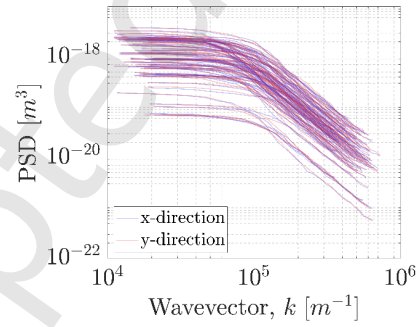
(a) Training samples



(b) Phase 1 testing samples



(c) Phase 2 testing samples



(d) Phase 3 testing samples

Figure 4 – Power spectral density (PSD) distributions of the rough surfaces used for training and testing.

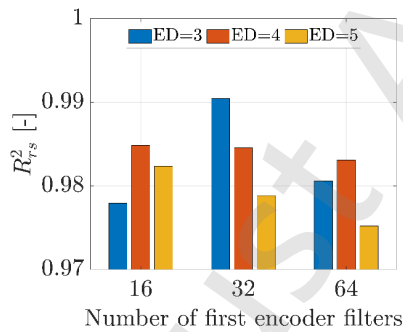
3.2. Neural network hyperparameters

The results of the Bayesian optimisation for the ANN hyperparameters are presented in Table 7.

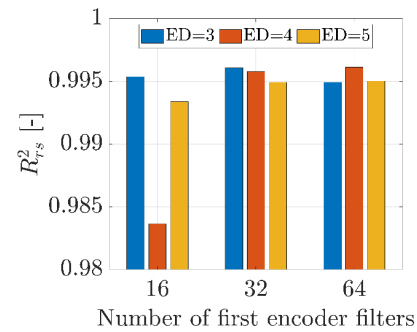
Table 7: Optimal ANN hyperparameters.

Parameter	ANN		
	Pressure	Film thickness	Temperature
Number of hidden layers	2	5	3
Number of neurons in each layer	256, 256	32, 512, 64, 128, 32	512, 64, 32
Dropout layer probability	0.1244	0.0510	0.0314

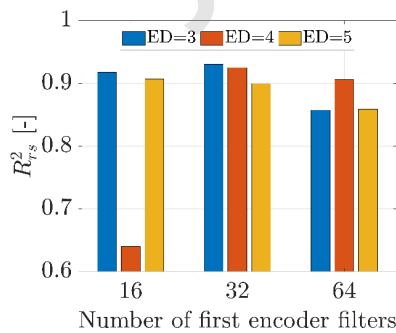
Figure 5 shows the effect the encoder depth (ED) and number of first encoder filters have on the mean accuracy of the U-Net predictions. The film thickness results are reported separately for the full profile and for the Hertzian region. This distinction is made because, as mentioned earlier, the macroscopic geometry component of the film thickness profiles was removed prior to training, allowing the network to focus on learning the roughness-induced variations within the contact region where predictive accuracy is most critical.



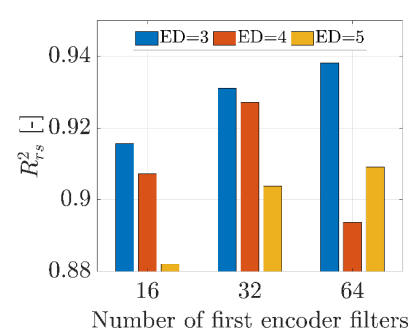
(a) Pressure



(b) Film thickness (full profile)



(c) Film thickness (Hertzian region)



(d) Temperature

Figure 5 – Effect of number of first encoder filters on U-Net mean prediction accuracy.

The results show that both the encoder depth and number of first encoder filters have some influence on the prediction accuracy, although that varies across the different fields. When using an encoder depth of 3, the optimal number of first encoder filters is 32 for pressure and for the Hertzian region film thickness results. For temperature, further increasing the number of first encoder filters to 64 results to a better mean prediction accuracy. This particular trend could be explained by the fact that the pressure and film thickness fields are less complex, being two-dimensional profiles, thus the prediction performance saturates once a sufficient number of first encoder filters is reached. However, the temperature fields require more first encoder filters due to their higher dimensionality, being three-dimensional profiles.

For an encoder depth of 4, the mean prediction accuracy for pressure is decreasing with an increasing number of first encoder filters, while for both Hertzian region film thickness and temperature results, the prediction accuracy is highest when the number first encoder filters is 32. Interestingly, there seems to be an outlier result for film thickness when the encoder depth is 4 and the number of first encoder filters is 16. This could be an artefact of the stochastic nature of neural network training and the imposing of early stopping of training to avoid overfitting of data, which occurred when the validation loss flatlined for 30 consecutive iterations in this study.

For an encoder depth of 5, the mean prediction accuracy for pressure and Hertzian region film thickness is decreasing with an increasing number of first encoder filters, while for the temperature field, the opposite is true as the mean prediction accuracy is increasing with an increasing number of first encoder filters.

Overall, an encoder depth of 3 seems to result to the most accurate predictions across all fields. With a 64×64 mesh, using an encoder depth of 4 or 5 generally results to lower prediction accuracy likely

due to excessive downsampling which reduces the spatial resolution required to resolve the localised asperity-induced features. Based on these results, an encoder depth of 3 was used for all three U-Nets, with the number of first encoder filters being 32 for the pressure and film thickness U-Nets and 64 for the temperature U-Net.

3.3. Accuracy of ANN predictions

Figure 6 shows the prediction accuracy of the ANNs, in terms of R_{rs}^2 , for smooth pressure, film thickness and temperature profiles. The mean R_{rs}^2 values are 0.998, 0.999 and 0.925 respectively. The corresponding mean root mean square error (RMSE) values are 0.014 GPa, 16.546 nm and 0.870 °C.

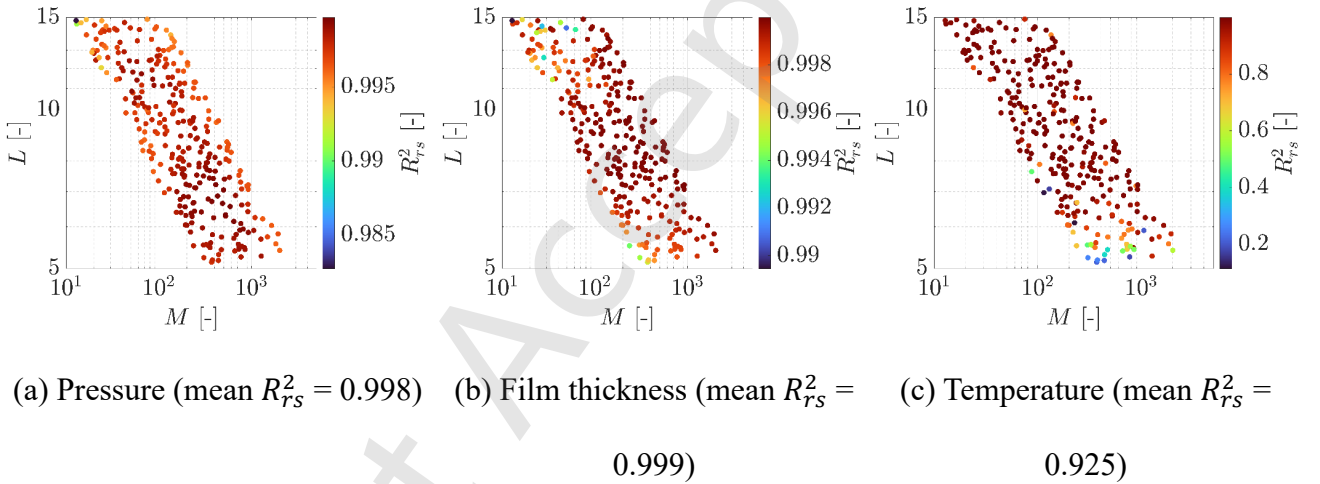


Figure 6 – Phase 1 ANN prediction accuracy.

The temperature predictions exhibit slightly reduced accuracy at lower values of L . Since lower L values correspond to lower entrainment speeds, the magnitude of the temperature rise in these cases is smaller. As highlighted in the previous study [12] there is a direct correlation between temperature profile magnitude and prediction accuracy, with higher magnitude profiles being implicitly prioritised during ANN training. Consequently, reduced temperature gradients at low L values lead to slightly lower predictive performance.

Compared to that previous study [12], the prediction accuracy of the ANNs is now higher for all three smooth profiles. Previously the reported R_{rs}^2 values were 0.997, 0.997 and 0.873 for pressure, film thickness and temperature, respectively. The observed improvement in the current work can be attributed to two primary factors. First, the number of input parameters has been reduced from twelve to three, simplifying the mapping between inputs and outputs. Second, the present dataset is restricted to the piezoviscous-elastic regime, whereas the earlier study included both piezoviscous-elastic and piezoviscous-rigid regimes. Limiting the parameter space to a single regime reduces variability in the underlying physical behaviour, thereby facilitating more accurate learning by the ANNs.

3.4. Accuracy of U-Net predictions

3.4.1 Phase 1

Figure 7 shows the prediction accuracy of the U-Nets for rough pressure, film thickness and temperature profiles across the operating parameter space. The mean R_{rs}^2 values are 0.990, 0.996 and 0.938 respectively with the corresponding mean RMSE values are 0.046 GPa, 42.971 nm and 1.510 °C.

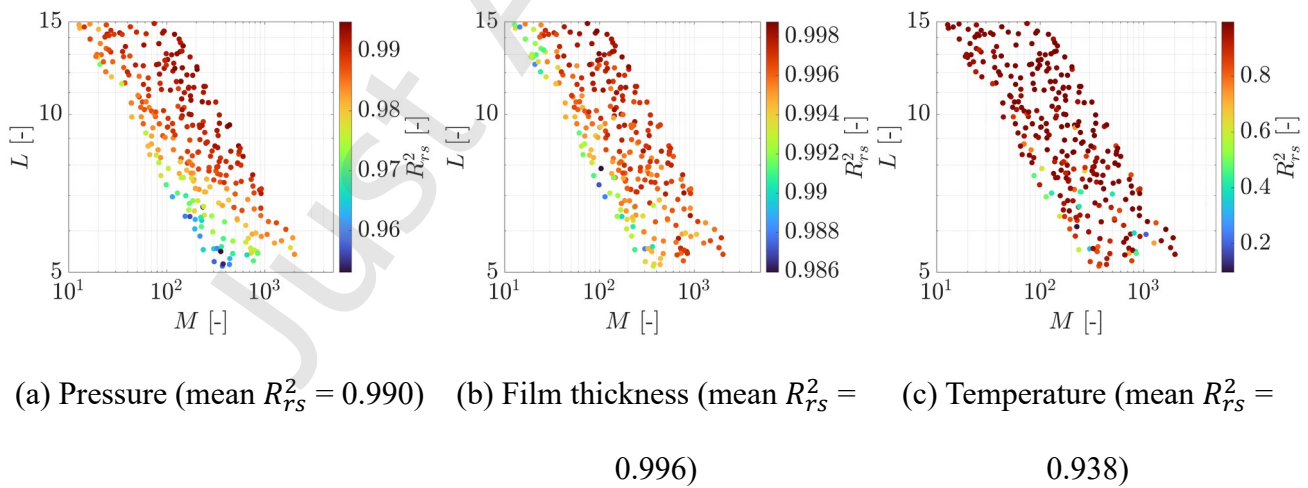


Figure 7 – Phase 1 U-Net prediction accuracy.

Despite the significantly increased spatial complexity of the rough profiles compared to the smooth solutions, the U-Nets achieve high predictive accuracy across all three output fields. This demonstrates

that the networks successfully learn the nonlinear mapping between the smooth solution, the imposed roughness distributions and the resulting rough mixed lubrication response.

Representative examples of classical mixed lubrication solver and U-Net-predicted rough profiles for pressure, film thickness and mid-film temperature are shown in Figure 8, Figure 9 and Figure 10 respectively. The predicted profiles closely reproduce both the global structure of the contact as well as the roughness-induced fluctuations. In particular, the location of local pressure peaks, film thickness perturbations and temperature variations are well captured, confirming that the U-Nets preserve both the macroscopic and microscopic features of the solution.

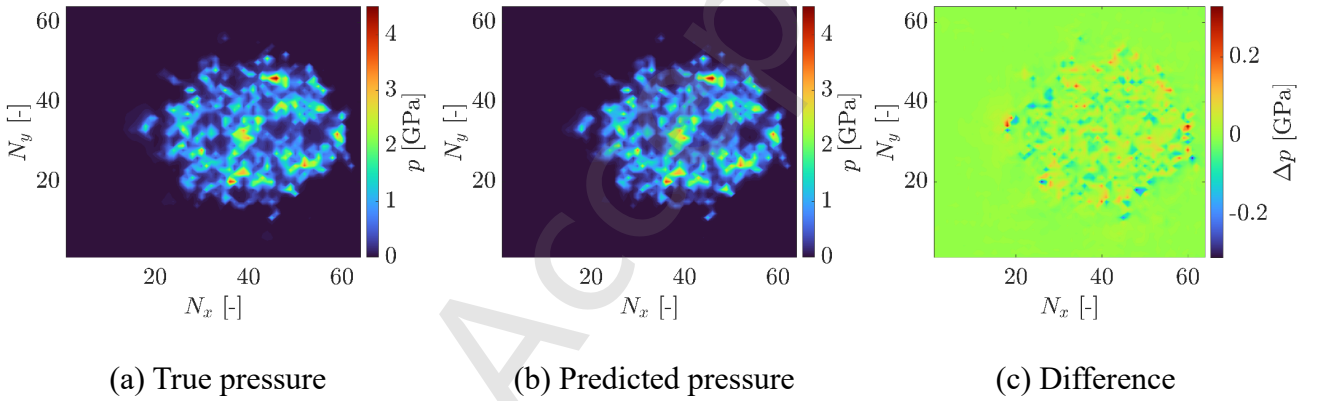


Figure 8 – Example U-Net pressure prediction ($R_{rs}^2 = 0.995$, $RMSE = 0.036$ GPa).

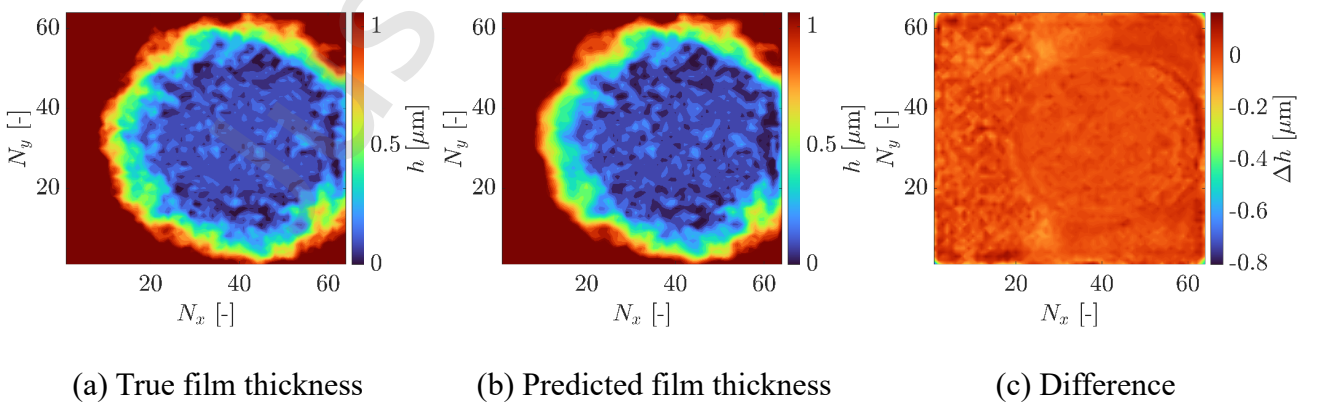
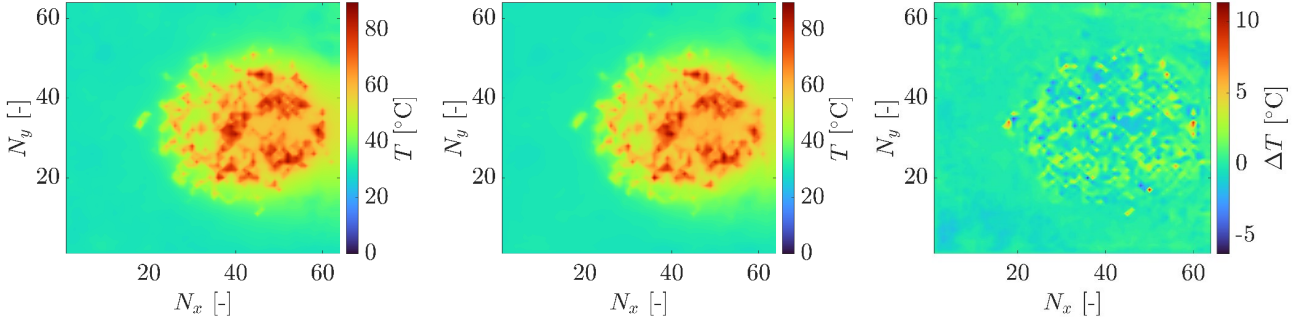


Figure 9 – Example U-Net film thickness prediction ($R_{rs}^2 = 0.997$, $RMSE = 46.254$ nm).



(a) True mid-film temperature

(b) Predicted mid-film

(c) Difference

temperature

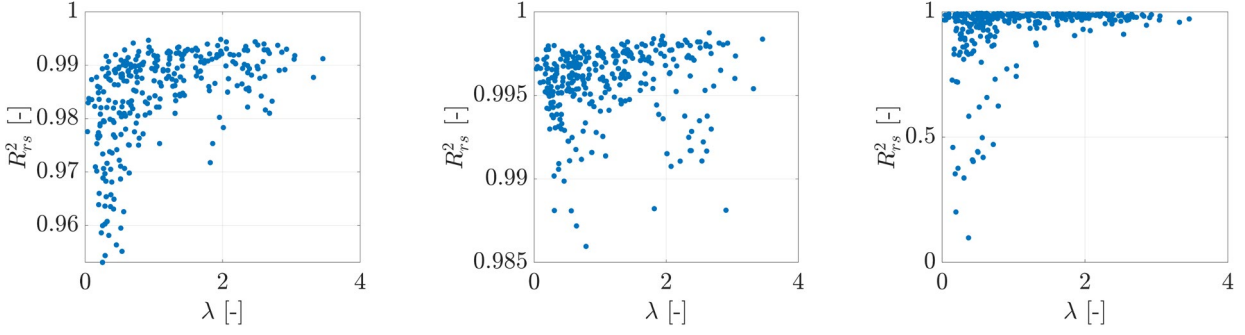
Figure 10 – Example U-Net mid-film temperature prediction ($R_{rs}^2 = 0.993$, $RMSE = 1.031$ °C).

As expected, the prediction accuracy for temperature is lower than for pressure and film thickness. The temperature field exhibits stronger localised gradients and is influenced by coupled thermal effects, making it inherently more difficult to reconstruct accurately. Nevertheless, the R_{rs}^2 value of 0.938 confirms that the multiscale structure of the rough temperature field is well captured. The overall performance in Phase 1 indicates that the U-Net architecture is capable of accurately reconstructing rough mixed lubrication fields when the testing samples are drawn from the same statistical roughness distribution as the training dataset.

Examining the distribution of R_{rs}^2 across the Moes $M - L$ operating parameter space in Figure 7 reveals distinct trends for each field. For pressure and film thickness, the lowest prediction accuracy occurs generally at lower M values, corresponding to lower loads, reduced elastic deformation, and less flattening of the roughness profile, hence more difficult conditions under which to accurately predict the resulting pressure and film shape. For pressure specifically, a reduction in accuracy is also observed at low L . Under these conditions, the pressure profile exhibits sharper gradients near the contact edges, more closely resembling the Hertzian dry contact solution in those regions compared to higher L , where hydrodynamic effects smooth out these gradients. This makes the resulting pressure

fields more difficult for the network to reconstruct accurately. In contrast, temperature accuracy shows no clear trend across the $M - L$ space, beyond a slight and comparatively inconsistent reduction at lower L values. This indicates that, for pressure and film thickness, prediction accuracy is primarily governed by the variability of the underlying field across the training dataset, whereas for temperature, prediction accuracy is instead governed by the presence and extent of asperity contact, given the strongly localised nature of contact-induced temperature peaks.

This distinction is further supported when prediction accuracy is instead examined as a function of the dimensionless film ratio, $\lambda = h_{ave}/R_q$, defined as the average film thickness in the Hertzian zone normalised by the synthesised RMS roughness of the upper and lower surfaces, which characterises the local lubrication regime. Figure 11 shows the relationship between R_{rs}^2 and λ for pressure, film thickness and temperature. A clear reduction in accuracy with decreasing λ is observed for pressure and temperature, whereas no clear trend is apparent for film thickness. For pressure, the reduction in accuracy at low λ is consistent with the low M and low L conditions discussed above, under which reduced elastic deformation and sharper, more Hertzian-like pressure gradients near the contact edges make the resulting pressure field inherently more difficult to reconstruct. Since the unified formulation of the hydrodynamic solver resolves hydrodynamic and asperity contact pressures within a single pressure field, this reduction in pressure accuracy is not attributable to the extent of asperity contact itself. In general, the pressure and film thickness predictions do not exhibit a specific sensitivity to the mixed lubrication regime itself. For temperature, however, the reduction in accuracy at low λ can be attributed to the increased prevalence of asperity contact under these conditions, since at contact locations, localised temperature peaks are disproportionately higher than in the surrounding non-contact regions, and the resulting sharp, spatially localised features are inherently more difficult for the network to reconstruct accurately.

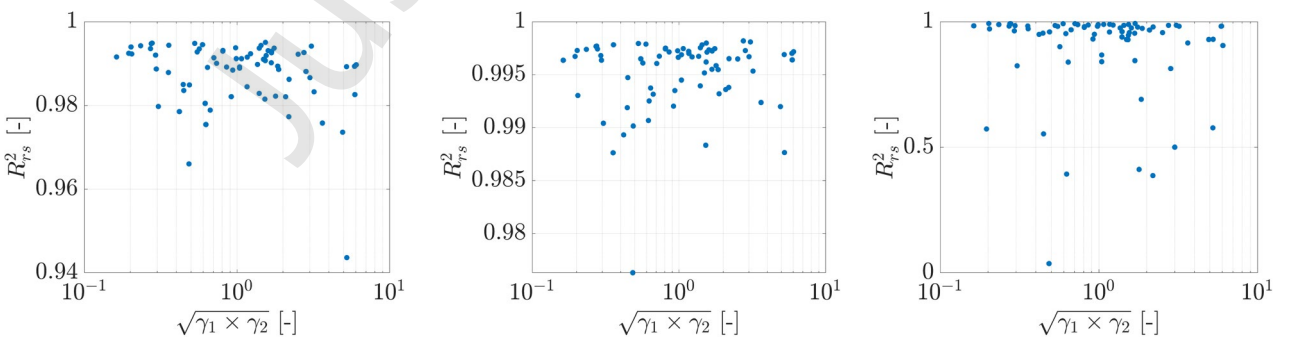


(a) Pressure (mean $R^2_{rs} = 0.990$) (b) Film thickness (mean $R^2_{rs} = 0.996$) (c) Temperature (mean $R^2_{rs} = 0.938$)

Figure 11 – Variation of Phase 1 U-Net prediction accuracy with film thickness parameter (λ).

3.4.2 Phase 2

Figure 12 shows the prediction accuracy of the U-Nets for rough pressure, film thickness, and temperature profiles as a function of the combined Peklenik number ($\gamma = \sqrt{\gamma_1 \times \gamma_2}$) of the top and bottom surfaces, which characterises surface anisotropy. The mean R^2_{rs} values are 0.988, 0.995 and 0.896 respectively with the corresponding mean RMSE values being 0.052 GPa, 47.546 nm and 1.657 °C. These results indicate that the networks maintain high predictive accuracy even when applied to surfaces with varying degrees of anisotropy, demonstrating good generalisation beyond the isotropic roughness distributions used during training.

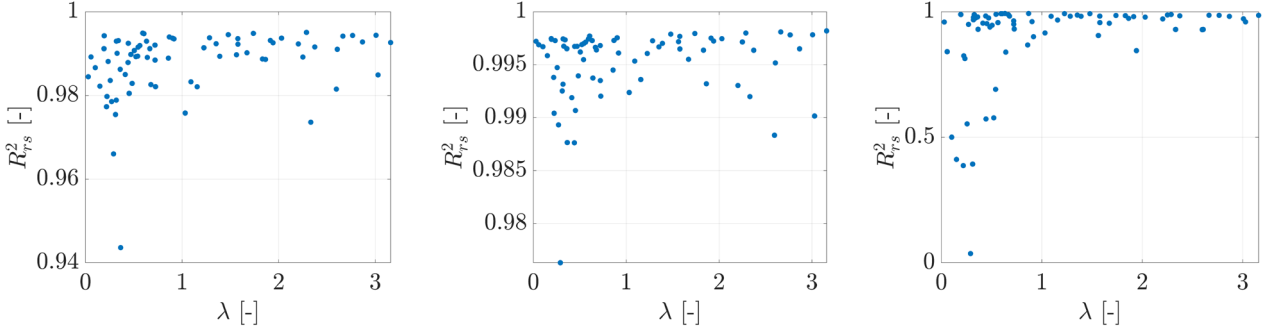


(a) Pressure (mean $R^2_{rs} = 0.988$) (b) Film thickness (mean $R^2_{rs} = 0.995$) (c) Temperature (mean $R^2_{rs} = 0.896$)

Figure 12 – Phase 2 U-Net prediction accuracy.

While no strong, systematic trend in prediction accuracy is observed as a function of γ for any of the three fields, the least accurate predictions do tend to occur as γ deviates further from the isotropic value ($\gamma = 1$) used in training. Pressure and film thickness accuracy values remain tightly clustered close to $R_{rs}^2 = 1$ across the entire γ range, with only a small number of scattered lower-accuracy samples. Temperature accuracy shows greater overall scatter, with several lower-accuracy outliers. These are more broadly distributed across the γ range, though again with a tendency towards reduced accuracy away from $\gamma = 1$. It cannot be entirely ruled out that transient effects introduced by the anisotropic roughness, such as sharper or more elongated local pressure and temperature features not present in the isotropic training data, could contribute to some of the reduction in accuracy in Phase 2. Another factor for the small decrease in accuracy could also be the smaller number of testing samples evaluated in Phase 2 (75 compared to 300 in Phase 1).

Where a clear trend is observed, once again, is with the dimensionless film ratio, $\lambda = h_{ave}/R_q$. As in Phase 1, prediction accuracy for all three fields was also examined as a function of λ . The same trend observed in Phase 1 is again apparent in Figure 13. The accuracy values for pressure and temperature decrease as λ decreases. For pressure, this is consistent with the sharper profiles that are more difficult to reconstruct as discussed in Phase 1. For temperature, the reduced accuracy at low λ reflects the increased prevalence and severity of asperity contact under these conditions, compounded in Phase 2 by the sharper, anisotropy-induced temperature peaks discussed above, both of which are inherently more difficult for the networks to reconstruct accurately.

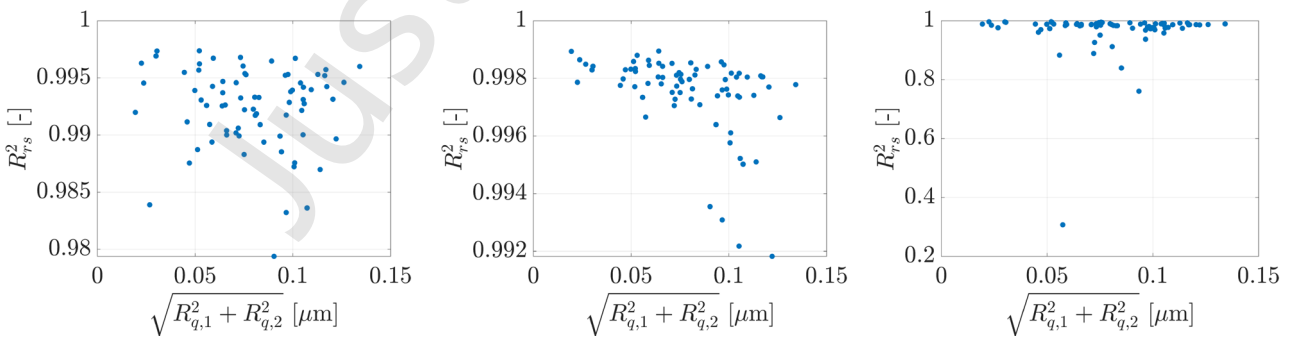


(a) Pressure (mean $R^2_{rs} = 0.988$) (b) Film thickness (mean $R^2_{rs} = 0.995$) (c) Temperature (mean $R^2_{rs} = 0.896$)

Figure 13 – Variation of Phase 2 U-Net prediction accuracy with film thickness parameter (λ).

3.4.3 Phase 3

Figure 14 shows the prediction accuracy of the U-Nets for rough pressure, film thickness, and temperature profiles as a function of the synthesised RMS roughness. In Phase 1, all surfaces had a fixed RMS of $0.1 \mu\text{m}$, whereas in Phase 3 the RMS was varied from $0.01 \mu\text{m}$ to $0.1 \mu\text{m}$, representing surfaces with different roughness amplitudes. The mean R^2_{rs} values across all tested RMS values are 0.992, 0.998, and 0.966 for pressure, film thickness, and temperature, respectively with corresponding mean RMSE values being 0.032 GPa, 35.203 nm, and 1.126°C .



(a) Pressure (mean $R^2_{rs} = 0.992$) (b) Film thickness (mean $R^2_{rs} = 0.998$) (c) Temperature (mean $R^2_{rs} = 0.966$)

Figure 14 – Phase 3 U-Net prediction accuracy.

These results demonstrate that the U-Nets can also generalise to surfaces with different RMS roughness values than those used during training, maintaining high predictive accuracy across all three output fields. In these figures, there is no clear trend of accuracy with RMS roughness. However, the least accurate film thickness predictions are occasionally found at higher RMS roughness values (right-hand side of the graph), suggesting that surfaces with larger roughness amplitudes can produce local features that are more challenging to reconstruct. Notably, the temperature predictions are slightly improved compared to Phase 1, likely because the inclusion of surfaces with smaller RMS roughness values introduces smoother local gradients, allowing the network to more accurately capture the multiscale thermal features.

As shown in Figure 15, the corresponding film ratio λ in Phase 3 extends up to approximately 10, which is substantially beyond the range of λ values encountered during training (approximately 0.03 to 3.5 in Phase 1). This provides a further test of the networks' ability to extrapolate to previously unseen lubrication regimes, in addition to the RMS roughness extrapolation already discussed above. The U-Nets maintain high prediction accuracy across this extended λ range for all three fields, confirming that the networks are able to generalise not only to previously unseen RMS roughness values, but also to film ratios well beyond those represented in the training dataset.

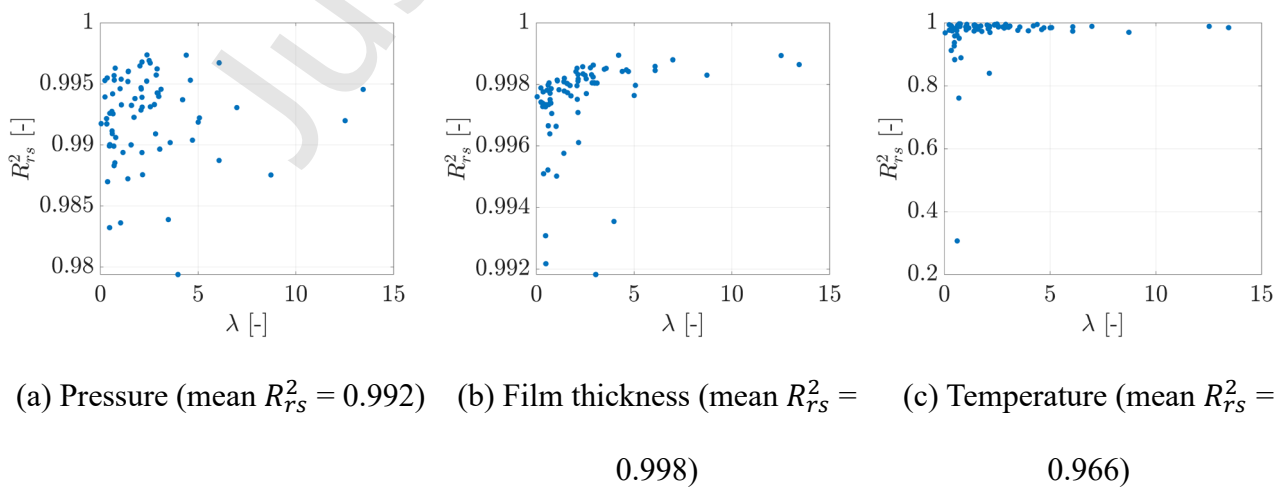


Figure 15 – Variation of Phase 3 U-Net prediction accuracy with film thickness parameter (λ).

Overall, Phases 2 and 3 demonstrate that the networks generalise well beyond the single, fixed roughness distribution used in training, accurately predicting rough lubrication fields for surfaces with previously unseen anisotropies and roughness amplitudes. This indicates that the networks do not need to be trained across the full space of surface characteristics to remain accurate, which has favourable implications for the practical training cost of the proposed approach, since fewer training samples, simulations and computational resources are required to develop the hybrid models.

3.5. Accuracy of hybrid framework solutions

Table 8 and Table 9 present a summary of the mean accuracy and RMSE values, respectively, of all the hybrid frameworks in Phase 1 of this study, with the Hybrid FW 3 values corresponding to the U-Net predictions discussed previously. It should be noted that the accuracy values for Hybrid FW 3 reported here may differ slightly from those presented earlier for the standalone U-Net predictions, since a small number of simulations (13 out of 300, a failure rate of 4.33%) initially failed to converge in hybrid frameworks incorporating ANN-predicted smooth fields. The underlying cause of such failures was traced to a bug in the classical solver, which has since been corrected. However, because this correction slightly altered the underlying solver physics and the neural networks were trained on data generated prior to the correction, the results for these 13 samples are excluded from the results below to ensure a fair comparison of framework performance under consistent training and evaluation conditions. This highlights an inherent limitation of data-driven machine learning models, which is that the neural network predictive accuracy fundamentally depends on the accuracy of the training data. If the underlying training data do change in any way, as occurred here following the solver correction, the models must be retrained to remain accurate.

Table 8: Mean accuracy (R_{rs}^2) of Phase 1 hybrid framework solutions.

Framework	p [-]	h [-]	T [-]	λ [-]	Contact area ratio [-]
Hybrid FW 1	1.000	1.000	0.999	1.000	1.000
Hybrid FW 2	1.000	1.000	0.986	1.000	1.000
Hybrid FW 3	0.990	0.995	0.939	0.998	0.631
Hybrid FW 1 (+ ANN)	1.000	1.000	0.978	1.000	0.999
Hybrid FW 2 (+ ANN)	1.000	1.000	0.965	1.000	0.999
Hybrid FW 3 (+ ANN)	0.990	0.995	0.921	0.998	0.663

Table 9: Mean RMSE of Phase 1 hybrid framework solutions.

Framework	p [GPa]	h [nm]	T [°C]	λ [-]	Contact area ratio [%]
Hybrid FW 1	0.001	0.135	0.162	0.000	0.042
Hybrid FW 2	0.007	0.467	0.573	0.001	0.056
Hybrid FW 3	0.046	44.981	1.538	0.032	4.040
Hybrid FW 1 (+ ANN)	0.002	0.740	0.315	0.007	0.190
Hybrid FW 2 (+ ANN)	0.007	0.883	0.712	0.007	0.190
Hybrid FW 3 (+ ANN)	0.046	45.275	1.582	0.034	3.773

Overall, the hybrid frameworks achieve extremely high levels of agreement with the classical solutions. The pressure and film thickness predictions remain almost perfectly accurate across all frameworks, with R_{rs}^2 values close to 1.000 and extremely small RMSE values. This indicates that the neural network predictions provide an excellent initial approximation for these fields, which are either preserved or only minimally corrected during the subsequent numerical iterations.

As expected, the highest level of agreement with the classical framework is obtained for Hybrid FW 1 and Hybrid FW 1 (+ANN), where both the smooth and rough iterations are allowed to converge fully. In these cases, the pressure and film thickness predictions are effectively identical to the classical solutions, while the temperature fields also maintain very high accuracy with very small RMSE values. In the previous study [12] it was shown that when the simulations were allowed to run to completion after the NN predictions, the converged solutions were identical to those obtained using the classical framework. In the present study, however, this is not strictly the case for the Hybrid FW 1 and Hybrid FW 1 (+ANN) solutions. These small discrepancies compared to the Classical FW are likely related to the less stringent temperature convergence criterion used in the current study (1×10^{-4} compared to 1×10^{-5} in [12]), which may allow small deviations introduced by the neural network predictions to persist in the final converged solution.

When the rough iteration is limited to a single global iteration in Hybrid FW 2 and Hybrid FW 2 (+ANN), the pressure and film thickness predictions remain largely unaffected, but a small reduction in temperature accuracy is observed. This is reflected in the reduced $R_{T_s}^2$ values and increased RMSE values for temperature compared to Hybrid FW 1, indicating that a single iteration is not always sufficient to fully correct the discrepancies present in the U-Net-predicted temperature field. The near-perfect accuracy of Hybrid FW 1 and 2 demonstrates that even less accurate initial U-Net predictions can be corrected once at least one global iteration of the classical solver is retained without introducing convergence issues.

The lowest accuracy is observed for Hybrid FW 3 and Hybrid FW 3 (+ANN), where the rough iteration is skipped entirely and the U-Net prediction is used directly as the final solution. While the pressure and film thickness predictions remain reasonably accurate, the temperature predictions exhibit a more

noticeable reduction in accuracy. This highlights the importance of the subsequent iterative solver acting as a correction step for refining the temperature NN predictions.

To further assess the accuracy of the hybrid frameworks in capturing asperity contact behaviour, Table 8 and Table 9 also report the accuracy of the predicted film ratio λ and contact area ratio. Hybrid FW 1 and Hybrid FW 2 (with and without ANN) predict contact area ratio with high accuracy ($R_{rs}^2 \geq 0.999$), consistent with their high film thickness accuracy. Hybrid FW 3 (with and without ANN), however, shows a substantially larger reduction in contact area ratio accuracy (R_{rs}^2 of 0.631 and 0.663, respectively) relative to its film thickness accuracy ($R_{rs}^2 \geq 0.995$ for both). This reduction arises because in the classical framework contact at a given node is defined by film thickness falling below a threshold value ($h/a \leq 1 \times 10^{-5}$), making the contact area ratio a binary, threshold-based metric that is inherently more sensitive to small film thickness deviations than the underlying continuous field itself. This reduction in contact area ratio accuracy does not translate into a corresponding reduction in pressure accuracy, since the unified formulation of the hydrodynamic solver resolves hydrodynamic and asperity contact pressures together within a single pressure field, without an explicit contact/non-contact distinction that could otherwise be affected by the wrong prediction of asperity contact locations. Temperature accuracy, however, is affected, since the governing thermal boundary conditions differ at asperity contact locations. Therefore, misclassification of contact at a given node introduces a corresponding error in the local temperature prediction, explaining the reduction in temperature accuracy observed for Hybrid FW 3. As such, the reduced contact area ratio accuracy for Hybrid FW 3 reflects this heightened sensitivity near the contact threshold, rather than a meaningful failure to capture the location or extent of asperity contact in the underlying film thickness field.

Overall, these results demonstrate that the hybrid frameworks can maintain extremely high predictive accuracy while reducing the computational effort associated with solving the rough mixed lubrication

problem. The classical iterative stages play an important role in correcting the neural network predictions, particularly for the temperature field, while the pressure and film thickness fields appear to be captured very accurately even with minimal iterative refinement.

3.6. Simulation time comparison

Table 10 summarises the mean simulation times obtained for all frameworks in Phase 1, together with the percentage reduction relative to the classical framework. The classical framework results to a mean simulation time of 10.36 minutes across all samples.

Table 10: Summary of Phase 1 simulation times.

Framework	Mean simulation time, t_{mean} [minutes]	Mean simulation time % difference compared to Classical FW
Classical FW	10.36	-
Hybrid FW 1	9.39	-9.36
Hybrid FW 2	7.89	-23.90
Hybrid FW 3	4.20	-59.43
Hybrid FW 1 (+ ANN)	6.98	-32.63
Hybrid FW 2 (+ ANN)	5.93	-42.78
Hybrid FW 3 (+ ANN)	2.65	-74.39

All hybrid frameworks provide a reduction in mean simulation time compared to the classical framework. The smallest improvement is observed for Hybrid FW 1, which reduces the mean simulation time by 9.36%. This relatively modest reduction is expected since the numerical solver is still allowed to converge fully during both the smooth and rough iterations.

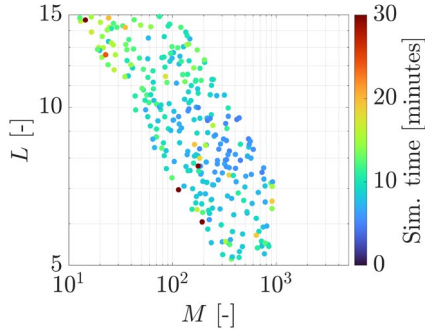
A larger reduction in simulation time is achieved for Hybrid FW 2, where the rough iteration is limited to a single global iteration. In this case, the mean simulation time decreases to 7.89 minutes, representing a 23.90% reduction. The reduced number of iterations limits the computational cost while maintaining high solution accuracy, as shown in the previous section.

The most significant improvement among the frameworks without ANN acceleration is obtained for Hybrid FW 3, where the rough iteration is skipped entirely and the U-Net prediction is used directly as the final rough solution. This reduces the mean simulation time to 4.20 minutes, representing a 59.43% decrease relative to the classical framework. This indicates that, on average, about 40% of the simulation time is spent on the smooth iteration, highlighting that the smooth solver accounts for a substantial portion of the total computational cost even when the rough iteration is fully bypassed.

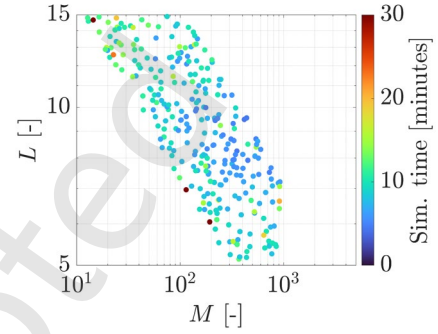
When the ANN predictions are incorporated into the framework, further reductions in simulation time are achieved. Hybrid FW 1 (+ANN) reduces the mean simulation time by 32.63%, while Hybrid FW 2 (+ANN) achieves a reduction of 42.78%. The largest decrease is obtained for Hybrid FW 3 (+ANN). In this case, the mean simulation time decreases to 2.65 minutes, corresponding to a reduction of 74.39% relative to the classical framework.

Figure 16 and Figure 17 provide a more detailed view of the simulation times across the Moes parameter space for hybrid frameworks without and with ANN acceleration, respectively. For the classical framework, simulation times generally increase with L , which is consistent with the previous study [12]. This trend arises because higher L values correspond to higher maximum temperatures in the contact, meaning the initial smooth solution is further from the final converged state, hence more iterations are required to reach convergence. Interestingly, a few simulations in the classical framework exhibit extreme simulation times which persist in Hybrid FW 1 and Hybrid FW 2. These outliers

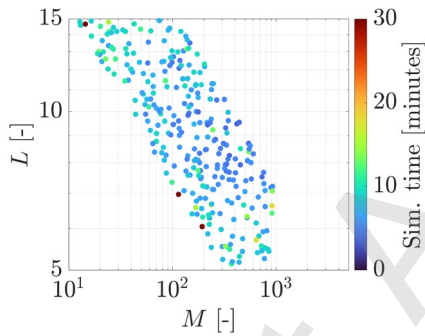
highlight cases where the rough iteration is particularly slow. In contrast, Hybrid FW 3 and Hybrid FW 3 (+ANN), which entirely replace the rough iteration with U-Net predictions, largely eliminate these extreme cases, demonstrating that the U-Net predictions effectively bypass the most computationally expensive stages of the solver.



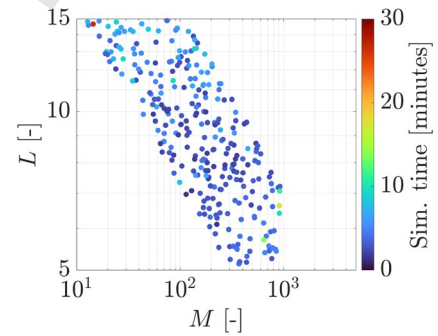
(a) Classical FW (mean $t = 10.36$ minutes)



(b) Hybrid FW 1 (mean $t = 9.39$ minutes)

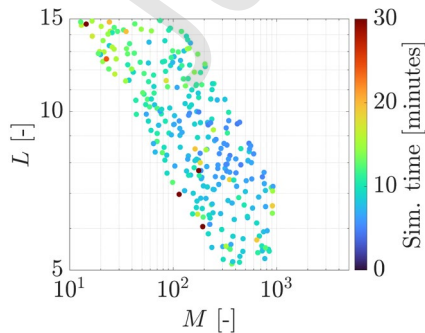


(c) Hybrid FW 2 (mean $t = 7.89$ minutes)

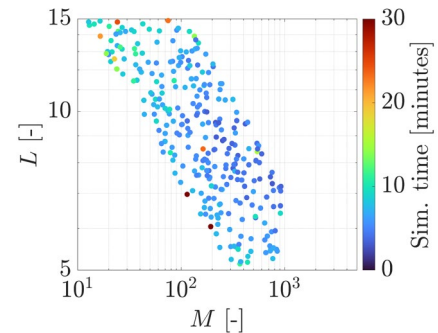


(d) Hybrid FW 3 (mean $t = 4.20$ minutes)

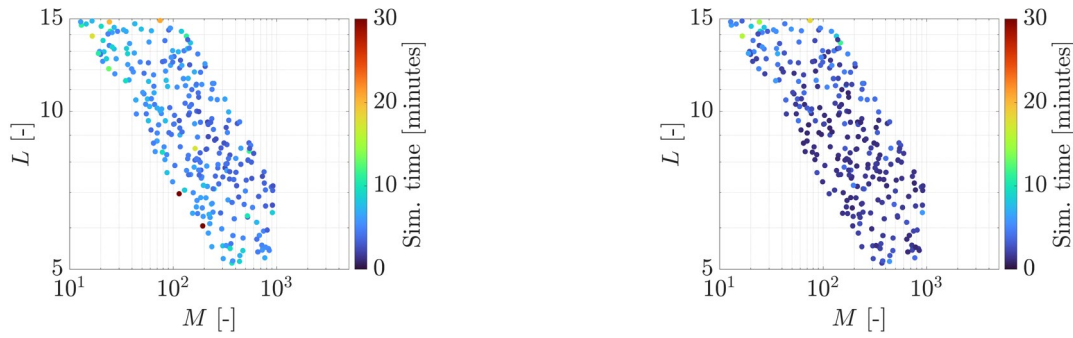
Figure 16 – Simulation times of Phase 1 hybrid frameworks without ANN acceleration.



(a) Classical FW (mean $t = 10.36$ minutes)



(b) Hybrid FW 1 (+ANN) (mean $t = 6.98$ minutes)



(c) Hybrid FW 2 (+ANN) (mean $t = 5.93$ minutes) (d) Hybrid FW 3 (+ANN) (mean $t = 2.65$ minutes)

Figure 17 – Simulation times of Phase 1 hybrid frameworks with ANN acceleration.

Overall, these results demonstrate that the hybrid frameworks can substantially reduce the computational cost of rough mixed lubrication simulations. The largest reductions are achieved when the neural network predictions replace the most computationally expensive stages of the numerical solver, with only a small impact on solution accuracy, particularly for the temperature field. It should be noted that the simulation times reported for these frameworks represent a “worst-case” scenario, as only two time iterations were performed in this study. In practical applications, where more rough iterations are typically required, the relative time savings for Hybrid FW 3 would likely be even greater.

In particular, conventional transient simulations typically introduce roughness gradually over many sequential time steps (e.g. 30-50, as discussed previously) before a converged, steady-state rough solution is reached. It is this steady-state solution that the proposed approach ultimately aims to predict. While the present study uses only two time iterations for both training and testing, in a practical application the networks would instead be trained on the steady-state rough solutions obtained from full transient simulations. The trained networks would then be able to predict the steady-state rough response directly from the smooth solution, replacing the many sequential time steps otherwise required to reach it with a single prediction step, and further increasing the time savings relative to

those reported in this study. Evaluating the hybrid frameworks' accuracy and acceleration when compared to full transient simulations thus remains an important direction for future work.

The computational savings achieved in this study also compare favourably with non-machine learning numerical acceleration techniques applied to related problems. For example, progressive mesh densification (PMD) has been reported to yield speed-ups of approximately 2–3 times relative to a classical fixed-mesh solver when applied to isothermal mixed EHL problems with deterministic surface roughness [35]. For smooth, isothermal cases, speed-ups of approximately 27 times have been reported for thick films, reducing to approximately 2 times for thinner films [36]. By comparison, the hybrid frameworks proposed in this study achieve reductions in mean simulation time of 74.39% (a speed-up of approximately 3.9 times) for rough-surface TEHL problems with individual simulation speed-ups of up to even 40 times, which is comparable to, or exceeds, the acceleration reported for PMD applied to deterministic rough EHL, while also incorporating coupled thermal effects not considered in the PMD studies cited.

It should be noted that these approaches are not mutually exclusive. Numerical acceleration methods such as PMD act on the discretisation and solution procedure of the classical solver itself, whereas the proposed hybrid framework replaces or reduces reliance on that solver via machine learning prediction. As such, PMD could, in principle, be applied within the classical solver component of the proposed hybrid frameworks to provide additional acceleration on top of that already achieved through the machine learning predictions.

In summary, the six hybrid frameworks presented in this study demonstrate that the degree of reliance on the classical iterative solver can be tuned to suit the accuracy and computational cost requirements of a given application, rather than identifying a single optimal configuration. Hybrid FW 1 (+ANN),

for instance, offers a particularly attractive balance of high accuracy and moderate acceleration, while Hybrid FW 3 (+ANN) offers substantially larger time savings at the cost of reduced temperature accuracy, which may be preferable in applications where some accuracy can be traded for greater speed, such as real-time monitoring or digital twins. The most appropriate framework is therefore expected to depend on the specific accuracy and speed requirements of the application in question.

Future work could focus on extending the framework to broader lubrication regimes and more complex, non-Gaussian surface topographies. In particular, improved treatment of thermal effects remains an important direction for enhancing predictive accuracy. This could be achieved through more optimised network architectures or by combining physics-informed training strategies with purely data-driven machine learning methods. While PINNs alone have shown limitations in capturing strongly nonlinear, tightly coupled TEHL behaviour, they may still offer value in ensuring physical consistency of the data-driven predictions.

4. Conclusions

This study has presented a hybrid computational framework for accelerating mixed lubrication simulations through the integration of ANNs and CNNs, specifically U-Nets, within a classical finite volume solver. The methodology extends previous work on smooth TEHL problems to deterministic rough surface conditions by enabling the direct prediction of rough pressure, film thickness and temperature fields.

A unified finite volume lubrication model was employed as the baseline solver. To address the high computational cost associated with rough surface simulations, particularly due to the gradual introduction of roughness over multiple time steps, a series of hybrid frameworks were developed. These frameworks combine ANN-based prediction of smooth fields and U-Net-based reconstruction

of rough fields, with varying levels of reliance on the classical iterative solver. This structure allows for a systematic trade-off between computational efficiency and solution accuracy.

The ANN models demonstrated high accuracy for smooth TEHL predictions, while the U-Net models achieved strong performance in predicting rough mixed lubrication fields, capturing both global contact behaviour and local roughness-induced perturbations. Pressure and film thickness were consistently predicted with very high accuracy across all test cases, while temperature remained the most challenging field due to its stronger nonlinear coupling and sharper spatial gradients.

Overall, the hybrid frameworks achieved substantial reductions in computational cost compared to the classical solver, with the most efficient configuration reducing simulation time by about 75% while maintaining good agreement with high-fidelity solutions. In addition, the hybrid approach mitigated extreme runtime variability associated with convergence issues in the classical solver, highlighting its potential for robust and efficient simulation.

These results demonstrate that hybrid machine learning–numerical approaches can provide a viable and efficient alternative to conventional mixed lubrication solvers, offering significant computational savings while preserving high predictive accuracy. As discussed previously, validating this capability for true transient rough-surface simulations using networks trained on steady-state data obtained from full transient simulations remains an important direction for future work. Once validated, this approach is expected to enable simulation capabilities suitable for real-time engineering applications such as digital twins and online condition monitoring systems.

CRedit authorship contribution statement

Conceptualization, F.K., S.A., D.D. and J.P.E.; Methodology, F.K. and S.A.; Software, F.K. and S.A.; Validation, F.K. and S.A.; Formal Analysis, F.K.; Investigation, F.K.; Resources, D.D. and J.P.E.; Data Curation, F.K.; Writing – Original Draft Preparation, F.K.; Writing – Review and Editing, S.A., D.D. and J.P.E.; Visualization, F.K.; Supervision, D.D. and J.P.E.; Project Administration, D.D.; Funding Acquisition, D.D. and J.P.E.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

F.K. thanks the UK Department of Science, Innovation and Technology (DSIT), the Engineering and Physical Sciences Research Council (EPSRC), and Shell for PhD funding through an iCASE studentship (EP/X524773/1). The authors thank Shell and the EPSRC for funding via the InFUSE Prosperity Partnership (EP/V038044/1). D.D. acknowledges the support of the Royal Academy of Engineering (RAEng) for the Shell/RAEng Research Chair in Complex Engineering Interfaces. J.P.E. acknowledges the support of the RAEng through their Research Fellowships scheme.

Data availability

Raw data used to plot the figures can be found online at: <https://doi.org/10.5281/zenodo.20491421> (to be activated on acceptance).

References

- [1] Q. J. Wang and D. Zhu, *Interfacial Mechanics*. First edition. | Boca Raton, FL : CRC Press/Taylor & Francis Group, 2019.: CRC Press, 2019. doi: 10.1201/9780429131011.
- [2] H. R. Braun, S. Korres, P. Laurs, and J. W. H. Franke, “Impact of ultra-low viscosity fluids on drivetrain functionality and durability,” *Lubricants*, vol. 9, no. 12, 2021, doi: 10.3390/lubricants9120119.
- [3] G. E. Morales-Espejel, “Surface roughness effects in elastohydrodynamic lubrication: A review with contributions,” 2014. doi: 10.1177/1350650113513572.
- [4] X. Ai and H. S. Cheng, “The Influence of Moving Dent on Point EHL Contacts,” *Tribology Transactions*, vol. 37, no. 2, pp. 323–335, Jan. 1994, doi: 10.1080/10402009408983301.
- [5] L. Chang, “A deterministic model for line-contact partial elastohydrodynamic lubrication,” *Tribol. Int.*, vol. 28, no. 2, pp. 75–84, Mar. 1995, doi: 10.1016/0301-679X(95)92697-4.
- [6] Y.-Z. Hu and D. Zhu, “A Full Numerical Solution to the Mixed Lubrication in Point Contacts,” *J. Tribol.*, vol. 122, no. 1, pp. 1–9, Jan. 2000, doi: 10.1115/1.555322.
- [7] D. Zhu and Y. Z. Hu, “A computer program package for the prediction of ehl and mixed lubrication characteristics, friction, subsurface stresses and flash temperatures based on measured 3-d surface roughness,” *Tribology Transactions*, vol. 44, no. 3, 2001, doi: 10.1080/10402000108982471.
- [8] C. H. Venner and A. A. Lubrecht, “Multilevel methods in lubrication,” *Tribology Series*, vol. 37, 2000, doi: 10.1243/1350650011543727.
- [9] S. Liu, Q. Wang, and G. Liu, “A versatile method of discrete convolution and FFT (DC-FFT) for contact analyses,” *Wear*, vol. 243, no. 1–2, pp. 101–111, Aug. 2000, doi: 10.1016/S0043-1648(00)00427-0.

- [10] S. Ardah, F. J. Profito, T. Reddyhoff, and D. Dini, “Advanced modelling of lubricated interfaces in general curvilinear grids,” *Tribol. Int.*, vol. 188, p. 108727, Oct. 2023, doi: 10.1016/j.triboint.2023.108727.
- [11] S. Ardah, F. J. Profito, and D. Dini, “An integrated finite volume framework for thermal elasto-hydrodynamic lubrication,” *Tribol. Int.*, vol. 177, p. 107935, Jan. 2023, doi: 10.1016/j.triboint.2022.107935.
- [12] F. Kaliafetis, S. Ardah, J. P. Ewen, and D. Dini, “Using artificial neural networks to accelerate thermo-elastohydrodynamic lubrication simulations,” *Tribol. Int.*, vol. 212, 2025, doi: 10.1016/j.triboint.2025.110978.
- [13] M. Marian and S. Tremmel, “Current trends and applications of machine learning in tribology—a review,” 2021. doi: 10.3390/LUBRICANTS9090086.
- [14] M. Marian, J. Mursak, M. Bartz, F. J. Profito, A. Rosenkranz, and S. Wartzack, “Predicting EHL film thickness parameters by machine learning approaches,” *Friction*, vol. 11, no. 6, pp. 992–1013, Jun. 2023, doi: 10.1007/s40544-022-0641-6.
- [15] W. Habchi and S. Bair, “Machine-Learning-Assisted Identification and Formulation of High-Pressure Lubricant-Piezoviscous-Response Parameters for Minimum Film Thickness Determination in Elastohydrodynamic Circular Contacts,” *Tribol. Lett.*, vol. 72, no. 4, p. 134, Dec. 2024, doi: 10.1007/s11249-024-01937-2.
- [16] J. Kelley, V. Schneider, G. Poll, and M. Marian, “Enhancing practical modeling: A neural network approach for locally-resolved prediction of elastohydrodynamic line contacts,” *Tribol. Int.*, vol. 199, 2024, doi: 10.1016/j.triboint.2024.109988.
- [17] C. Zhang, M. Tošić, J. Kelley, W. Habchi, M. Marian, and T. Lohner, “A Modular and Generalized Machine Learning Approach for Non-Newtonian EHL Film Thickness Prediction,” *Tribol. Int.*, p. 112211, May 2026, doi: 10.1016/j.triboint.2026.112211.

- [18] M. Raissi, P. Perdikaris, and G. E. Karniadakis, “Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations,” *J. Comput. Phys.*, vol. 378, 2019, doi: 10.1016/j.jcp.2018.10.045.
- [19] A. Almqvist, “Fundamentals of Physics-Informed Neural Networks Applied to Solve the Reynolds Boundary Value Problem,” *Lubricants*, vol. 9, no. 8, p. 82, Aug. 2021, doi: 10.3390/lubricants9080082.
- [20] M. Marian and S. Tremmel, “Physics-Informed Machine Learning—An Emerging Trend in Tribology,” *Lubricants*, vol. 11, no. 11, p. 463, Oct. 2023, doi: 10.3390/lubricants11110463.
- [21] S. Cuomo, V. S. Di Cola, F. Giampaolo, G. Rozza, M. Raissi, and F. Piccialli, “Scientific Machine Learning Through Physics-Informed Neural Networks: Where we are and What’s Next,” *J. Sci. Comput.*, vol. 92, 2022, doi: 10.1007/s10915-022-01939-z.
- [22] S. Zhao, L. Hou, R. Yang, and G. Zhou, “An adaptively weighted multi-scale physics-informed neural network for elastohydrodynamic lubrication in rough-surface point contacts,” *Tribol. Int.*, vol. 224, p. 112270, Dec. 2026, doi: 10.1016/j.triboint.2026.112270.
- [23] J. Gu *et al.*, “Recent advances in convolutional neural networks,” *Pattern Recognit.*, vol. 77, 2018, doi: 10.1016/j.patcog.2017.10.013.
- [24] W. Xu *et al.*, “High-resolution u-net: Preserving image details for cultivated land extraction,” *Sensors (Switzerland)*, vol. 20, no. 15, 2020, doi: 10.3390/s20154064.
- [25] W. Yang, R. Zhang, S. K. K. Chow, and Z. Li, “A survey of U-Net variant network for MRI brain tumor segmentation,” *Discover Artificial Intelligence*, 2025, doi: 10.1007/s44163-025-00525-0.

- [26] C. Bian and X. Chen, “MambaClinix: hierarchical gated convolution and mamba-based u-net for enhanced 3D medical image segmentation,” *Biomed. Signal Process. Control*, vol. 124, p. 110657, Sep. 2026, doi: 10.1016/j.bspc.2026.110657.
- [27] B. Feng, S. Lyu, and Y. He, “Brain tumour segmentation using learnable wavelet transform and attention mechanism network,” *Biomed. Signal Process. Control*, vol. 122, p. 110444, Aug. 2026, doi: 10.1016/j.bspc.2026.110444.
- [28] Y. Matsuzaka and M. Iyoda, “Applications, image analysis, and interpretation of computer vision in medical imaging,” *Frontiers in Radiology*, vol. 5, 2026, doi: 10.3389/fradi.2025.1733003.
- [29] F. Kaliafetis, D. Dini, J. P. Ewen, and S. Ardah, “A Finite Volume-Based Unified Transient Deterministic Framework for Lubrication Modelling,” *Lubricants*, vol. 14, no. 7, p. 281, Jul. 2026, doi: 10.3390/lubricants14070281.
- [30] PCS Instruments, “MTM.” Accessed: Apr. 07, 2025. [Online]. Available: <https://pcs-instruments.com/product/mtm/>
- [31] A. Pradhan *et al.*, “The Surface-Topography Challenge: A Multi-Laboratory Benchmark Study to Advance the Characterization of Topography,” *Tribol. Lett.*, vol. 73, no. 3, 2025, doi: 10.1007/s11249-025-02014-y.
- [32] T. Agrawal, *Hyperparameter Optimization in Machine Learning*. Berkeley, CA: Apress, 2021. doi: 10.1007/978-1-4842-6579-6.
- [33] H. Moes, “Optimum similarity analysis with applications to elastohydrodynamic lubrication,” *Wear*, vol. 159, no. 1, pp. 57–66, Nov. 1992, doi: 10.1016/0043-1648(92)90286-H.
- [34] G. Nijenbanning, C. H. Venner, and H. Moes, “Film thickness in elastohydrodynamically lubricated elliptic contacts,” *Wear*, vol. 176, no. 2, pp. 217–229, Aug. 1994, doi: 10.1016/0043-1648(94)90150-3.

- [35] W. Pu, J. Wang, and D. Zhu, “Progressive Mesh Densification Method for Numerical Solution of Mixed Elastohydrodynamic Lubrication,” *J. Tribol.*, vol. 138, no. 2, Apr. 2016, doi: 10.1115/1.4031495.
- [36] D. Zhu, “Progressive Mesh Densification (PMD) Method,” in *Encyclopedia of Tribology*, Boston, MA: Springer US, 2013, pp. 2690–2696. doi: 10.1007/978-0-387-92897-5_635.

Just Accepted